

Evaluating Generative AI as a Mentor Resource: Bias and Implementation Challenges

Jimin Lee* and Alena G. Esposito

Department of Psychology, Clark University

Worcester, MA, USA

jmllee@clarku.edu aesposito@clarku.edu

Abstract

This study explored how students' perceptions of helpfulness and caring skew their ability to identify AI versus human mentorship responses. Emotionally resonant responses often lead to misattributions, indicating perceptual biases that shape mentorship judgments. The findings inform ethical, relational, and effective integration of AI in student support.

1 Introduction

Mentorship in higher education is widely recognized as a developmental relationship in which mentors offer academic, psychosocial, and emotional guidance to support students' success and growth (Nuis et al., 2023). Through sharing expertise and personal experience, mentors help students expand their knowledge base and pursue individual goals (Köbis and Mehner, 2021). As Generative Artificial Intelligence (GenAI) tools become increasingly integrated into academic settings, their role is expanding beyond academic support and research assistance to include potential contributions to mentoring relationships.

Upon entering college, students often encounter a combination of formal mentorship, typically through faculty advisors, and informal mentoring through peers or other institutional contacts (Rhodes et al., 2000; Jacobi, 1991). Understanding how these relationships form and function is critical to fostering positive and developmental outcomes. The rise of GenAI tools, such as ChatGPT (OpenAI, 2024a), prompts renewed reflection on how students engage with mentoring and what constitutes meaningful support in both human and machine-mediated contexts. Early evidence suggests that GenAI may function as a mentoring-like resource, offering students guidance and feedback that mimics the conversational tone of a human

tutor (Le et al., 2025; Javaid et al., 2023). This insight highlights the need to examine how and why students may turn to GenAI for informal support and guidance.

GenAI tools can serve not only as tutors but also as supportive companions, helping reduce feelings of isolation and disconnection in academic environments (Farrelly and Baker, 2023). This growing interest in AI as a mentor-like resource is also shaped by broader concerns about burnout and mental health in higher education, which affect not only students but also faculty mentors who must balance teaching, research, and administrative demands (Hammoudi Halat et al., 2023). As institutions seek solutions to these overlapping pressures, GenAI presents both opportunities and challenges.

While GenAI facilitates academic learning by assisting with writing, problem-solving, and research tasks (Baidoo-Anu and Owusu Ansah, 2023; Le et al., 2025; Montenegro-Rueda et al., 2023; Schönberger, 2023), it still lacks the nuanced relational and developmental depth of human mentorship (Dempere et al., 2023). Ethical concerns and AI literacy are essential components of its responsible implementation, but so too is understanding students' lived perceptions of these tools. For GenAI to be effectively integrated into mentorship, educators and AI designers must understand how students evaluate its usefulness and trustworthiness. This factor is especially important in light of evidence that AI systems can unintentionally amplify human biases, especially in emotionally or socially sensitive domains, and that users may not always be aware of AI's influence on their perceptions and judgments (Glickman and Sharot, 2025).

Our prior work has explored these questions by examining how students interpret and engage with both AI-generated and human-authored responses in simulated mentorship scenarios. Drawing on the Perceptual Bias Activation (PBA) framework (Lee and Esposito, 2025b), we investigated whether stu-

*Corresponding author.

dents' evaluations of response quality and accuracy of source identification were shaped by cognitive biases when the authorship of response sources differed across contexts, with the lowest accuracy in personal, mental-health-related scenarios. This finding may suggest that GenAI tools blend seamlessly into mentorship roles in mental health contexts but also raise concerns about overreliance on AI. Follow-up analyses in the personal domain further demonstrated that responses perceived as AI-authored were consistently rated as less helpful and caring, regardless of their actual source. However, when examined by actual authorship, AI-generated responses were rated as more caring than human responses. To further explore this discrepancy, we conducted an Inductive Content Analysis (ICA) of participants' open-ended explanations (Lee et al., 2025). The analysis revealed that source attributions were influenced by features such as tone, language, and perceived emotional depth, highlighting that students' interpretations were guided more by their perceptions and assumptions than by the intrinsic qualities of the response, which points to a lack of familiarity with GenAI tools for mental health support.

We also examined individual-level factors that might influence source accuracy and evaluation in all domains (Lee and Esposito, 2025a). Prior experience using GenAI was positively associated with more accurate source identification, suggesting that familiarity with GenAI tools may reduce perceptual bias. On the other hand, students' mentorship background (e.g., having a faculty mentor, peer mentor, or mental health counselor) did not predict improved source recognition. Using the Unified Theory of Acceptance and Use of Technology (UTAUT; (Venkatesh et al., 2003)), we found that students who rated GenAI responses as more useful, easier to use, and socially acceptable were more likely to evaluate them favorably, but only when they believed the response was AI-generated. These findings point to the need for greater transparency and intentional AI literacy efforts within higher education.

Our prior work reveals how perceptual biases can influence students' engagement with GenAI tools, often leading them to undervalue these resources, including in situations where the information is more readily available than from a human mentor. This raises two key questions: To what extent do perceptual biases limit the integration of Generative AI as a mentorship resource? And what factors, if

any, mitigate this bias?

The current study seeks to address these two questions by reversing the analytical lens. Instead of examining how perceived or actual authorship affects evaluations, we ask: Are students more accurate in identifying the source of mentorship responses when they find those responses more helpful or caring? In other words, do positive evaluations enhance or cloud students' source discernment? We combine quantitative and qualitative analyses to explore this question. Specifically, we investigate whether students' ratings of helpfulness and caring predict their accuracy in identifying response sources, and how these patterns differ across personal, social, and academic mentorship contexts. We also analyze open-ended explanations from students to better understand the features that inform their judgments. This mixed-methods approach deepens our understanding of how perceptual biases shape students' interactions with human and AI mentorship. Furthermore, the findings of this study will have critical implications for the design and implementation of GenAI in higher education, particularly as institutions seek to balance technological innovation with relational and developmental support for students.

2 Methods

2.1 Participants

Our dataset stems from a larger project (Lee and Esposito, 2025a; Lee et al., 2025) that explored students' perceptions of GenAI and faculty mentorship in higher education. The study received approval from the college's Institutional Review Board (IRB Protocol #546). Although these data have previously been analyzed and published, the current study addresses new research questions and employs extended analytical approaches.

A total of 147 undergraduate students ($M_{age} = 19.34$ years, $SD_{age} = 1.33$ years, 105 female, 37 male, 2 non-binary, and 3 prefer not to answer) were recruited from a small liberal arts college in the northeastern United States. The sample was racially and ethnically diverse: 14. 97% Asian, 6. 80% Black, 67. 35% White, and 10. 88% Hispanic. All participants were at least 18 years old and provided informed consent prior to participation.

2.2 Procedure

The secure Qualtrics survey, which took approximately 30 minutes to complete, began with de-

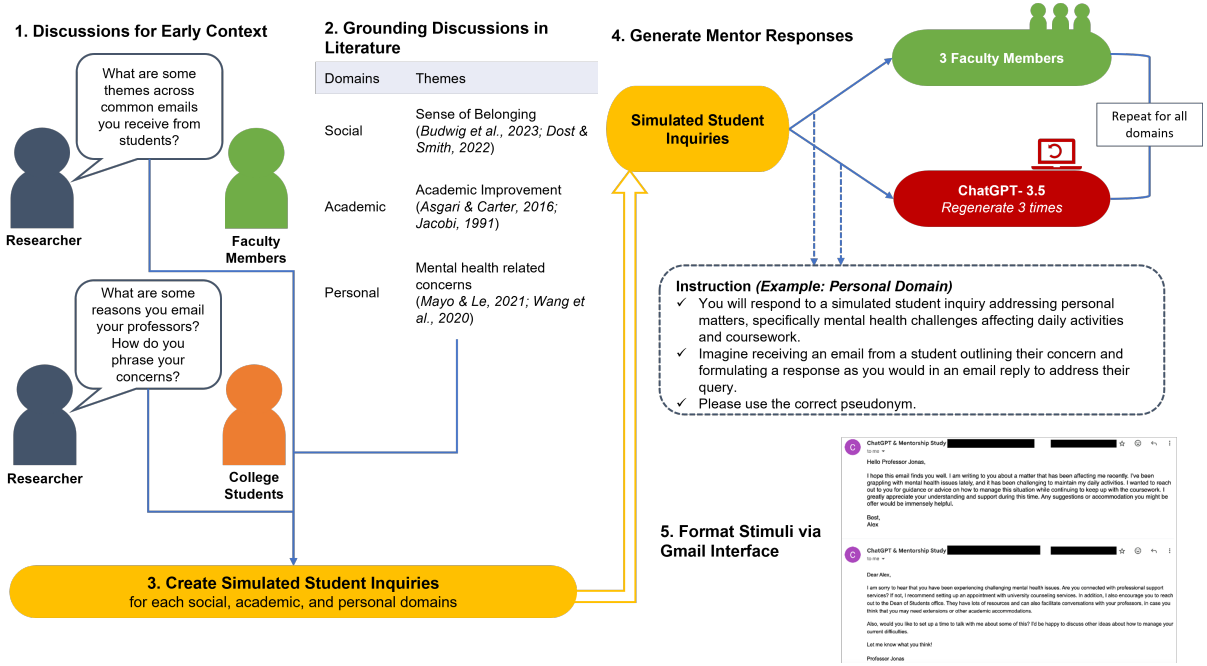


Figure 1: Contextualizing and creating stimuli.

mographic questions, followed by participant evaluations of mentorship interactions. Among the scenarios, the personal domain focused on mental health-related issues (Mayo and Le, 2021; Wang et al., 2020), the social domain focused on sense of belonging (Budwig et al., 2023; Dost and Mazzoli Smith, 2023), and the academic domain focused on academic improvement (Asgari and Carter, 2016; Jacobi, 1991). These domains were selected to reflect a realistic and broadly relevant context in which students seek support from their mentors (see Figure 1).

To explore perceptions of AI-generated versus human responses, participants were presented with three randomized and masked responses drawn from a pool of 18 responses (nine from ChatGPT version 3.5 (OpenAI, 2024b) and nine from human faculty members from three different academic disciplines who had received institutional awards or recognition for mentorship excellence within the past five years). Both ChatGPT and human faculty received identical prompts simulating student inquiries.

For the AI-generated responses, we regenerated three responses for each domain to maintain parity across conditions. Faculty members provided their responses based on previous mentoring experiences and did not use GenAI tools in drafting their replies. All responses were then reformatted to resemble the Gmail interface, reflecting the stan-

dard communication format used in many higher education settings.

Participants were instructed to identify whether each response was AI- or human-generated, without receiving feedback on their accuracy (see Table 1). This identification task was designed to activate perceptual biases. Once participants formed an impression of the source, this initial judgment could influence their subsequent evaluation of the response's quality and characteristics.

	AI	Human
Domain	%	%
Social	75.81	75.81
Academic	72.03	74.14
Personal	56.57	76.76

Table 1: Accuracy percentage of AI and human responses by domain.

After each identification, participants rated the response on a 5-point Likert scale (1= Not at all, 5= Extremely) across dimensions of helpfulness and caring (see Table 2).

They also provided written explanations for why they believed the response was from AI or a human, and why they rated it as they did, which served as our qualitative data. Following this evaluation task, participants completed additional survey measures assessing their broader perceptions of mentorship and AI in academic contexts.

Domain	Scales	Perceived Source				Actual Source			
		AI		Human		AI		Human	
		Mean	SD	Mean	SD	Mean	SD	Mean	SD
Social	Helpful	2.90	1.04	3.92	0.78	3.13	1.03	3.67	1.00
	Caring	2.67	1.01	4.00	0.85	2.91	1.12	3.74	1.02
Academic	Helpful	2.91	1.06	3.66	1.03	3.19	1.07	3.40	1.14
	Caring	2.73	1.08	3.56	1.16	3.06	1.17	3.25	1.21
Personal	Helpful	3.04	1.03	3.82	0.99	3.52	1.11	3.48	1.04
	Caring	2.93	1.03	3.77	1.06	3.54	1.11	3.33	1.14

Table 2: Ratings of helpfulness and caring by domain for perceived and actual source. *SD*= standard deviation.

3 Method of Analysis

To address our research question, we used both quantitative and qualitative measures. Quantitative data were analyzed using R (R version 4.3.2, R version 4.4.2) and RStudio (RStudio version 2024.09.1+394, RStudio version 2024.12.1+563) (R Core Team, 2024). To examine whether the helpfulness (Model 1) and care (Model 2) ratings predict the accuracy of the source across domains, we first performed binary logistic linear regression analyses. The reference category for our domain variable was set to Personal, where ratings were consistently higher across all scales.

We then used Inductive Content Analysis (ICA), a qualitative method used to identify patterns in textual data and support exploratory findings. ICA is particularly appropriate in contexts where prior research is limited, as it allows researchers to derive insights directly from the data through systematic coding and theme identification (Vears and Gillam, 2022). Given its applicability to various forms of written text, ICA was especially suitable for our study’s purpose of exploring human perceptions and experiences, independent of the specific mode of data collection (Elo and Kyngäs, 2008). The final thematic structure consisted of six overarching categories and 17 subthemes (Table 3).

Using the finalized codebook, we independently coded the qualitative responses. We have previously presented partial results in the personal domain (Lee et al., 2025), but we extended the ICA coding to include the social and academic domains for this study. To ensure analytic consistency and rigor, coding discrepancies were reviewed through a collaborative resolution process (Kyngäs, 2020). When disagreements arose, we held structured consensus-building sessions in which coders explained their rationale for coding decisions (Forman and Damschroder, 2008). Final coding decisions were reached through negotiated agreement.

Main Category	Generic Categories	Sub-Categories
Students’ Perceptions of Human vs. AI Mentorship	Tone of Response	Sincerity & Empathy Warmth & Approachability Professionalism & Formality
	Language	Authentic & Natural Clarity & Simplicity Structure & Format
	Information and Resources	Specific Information Resource Guidance Campus Knowledge
	Individualized Support & Actionable Advice	Personalized & Applicable Contextualized Understanding & Support
	Personal Connection	Emotional Connection Establishing Direct Connection in Person Genuine Investment in Student
	Holistic Student Support	Sense of Support Mental Health Overall Well-being and Growth

Table 3: Codebook developed and used for inductive content analysis.

4 Results

We investigated whether students’ ratings of Helpfulness (Model 1) and Caring (Model 2) predict their accuracy across the domains.

4.1 Helpfulness and Domain Predicting Accuracy

A binary logistic regression was conducted to examine whether students’ ratings of helpfulness predicted their ability to accurately identify the source of mentorship responses (human vs. AI) and whether this relationship differed across personal, social, and academic domains (see Table 4).

Predictors	Accuracy			
	Odds Ratios	SE	CI	p
(Intercept)	4.53	1.68	2.22–9.52	<.001*
Helpfulness	0.79	0.08	0.65–0.96	.019*
Domain [Social]	0.46	0.24	0.16–1.28	.137
Domain [Academic]	0.64	0.32	0.24–1.71	.374
Helpfulness × Domain [Social]	1.45	0.21	1.09–1.94	.011*
Helpfulness × Domain [Academic]	1.24	0.17	0.94–1.63	.126
Observations	1289			
Tjur’s R^2	0.015			

Table 4: Accuracy predicted by helpfulness and domain. *Indicates $p < .05$; SE = standard error; CI = 95% confidence interval.

Model 1 revealed a significant main effect of Helpfulness ($OR= 0.79$, $p = 0.019$, 95%CI [0.65, 0.96]), suggesting that higher helpfulness ratings were associated with lower odds of correctly identifying the response source. There was no significant main effect of domain (Social: $p = .137$; Academic: $p = .374$).

Qualitative responses indicated that when students misattributed authorship, it was due to tone, language, personalization, and informativeness. For example, a human-written response was misidentified as AI because it felt “too formal and robotic” (Tone of Response, P67), while another participant described a different human-generated response as “pretty basic without many specific details” (Information and Resource, P68). Several AI-generated responses were perceived as human due to emotionally resonant or personalized phrasing, such as showing “genuine appreciation and understanding of students’ struggles” (Tone of Response, P55). These patterns highlight how AI responses can be anthropomorphized, while human mentors may also provide rigid or impersonal responses that fail to meet students’ relational expectations.

Furthermore, there was a significant interaction between Helpfulness and the Social domain, indicating that in contexts related to sense of belonging, higher helpfulness ratings were positively associated with identification accuracy. Qualitative insights help explain this result. In the Social domain, participants were more likely to correctly identify human responses when they involved personal outreach, such as “offers to talk with students personally to give suggestions” (Personalized Guidance, P32), or when the tone conveyed compassion while respecting autonomy (“compassionate yet prioritizes the students’ autonomy, privacy, and space,” Language, P4). In contrast, responses perceived as checklist-like or impersonal were correctly identified as AI, as in comments like “feels incredibly impersonal and provides a checklist more than someone trying to communicate” (Language, P111) or “the advice would work for any university” (Personalized Guidance, P135). The interaction between Helpfulness and the Academic domain was not statistically significant ($p = .126$).

4.2 Caring and Domain Predicting Accuracy

A second logistic regression tested whether perceived Caring ratings predicted source identification accuracy, and whether this relationship varied across domains (see Table 5).

Predictors	Accuracy			
	Odds Ratios	SE	CI	p
(Intercept)	3.83	1.32	1.98–7.62	<.001*
Caring	0.83	0.08	0.69–0.99	.043*
Domain [Social]	0.72	0.35	0.28–1.88	.501
Domain [Academic]	1.46	0.69	0.58–3.71	.427
Caring×Domain [Social]	1.28	0.17	0.98–1.67	.073
Caring×Domain [Academic]	0.97	0.13	0.75–1.25	.794
Observations	1289			
Tjur’s R^2	0.017			

Table 5: Accuracy predicted by caring and domain. *Indicates $p < .05$; SE = standard error; CI = 95% confidence interval.

The results showed a significant main effect of Caring, indicating that higher caring ratings were also associated with lower odds of accurate source identification. Though domain effects were not significant (Social: $p = .501$; Academic: $p = .427$), nor were the interactions between Caring and Domain (Social: $p = .073$; Academic: $p = .794$), our qualitative data illustrate perceptual bias towards responses.

Participants interpreted emotionally validating or well-phrased AI responses as human-authored. Participants reported “[the response] indicated the importance of our well-being” (Holistic Student Support, P58) and “used thoughtfully placed words to show validation and support” (Language, P43). These examples illustrate how AI’s capacity to mimic affective tone can lead to over-attribution of caring intent and misidentification. Conversely, human responses perceived as distant or overly formal were misattributed as AI. One participant stated the response “felt a bit cold” (Tone of Response, P64), while another described it as a “scripted response” (Language, P77). Even when human mentors intended to convey care, lack of emotional language or concrete support diminished perceived authenticity: “appears to want to be supportive but does not provide the support in any tangible way” (Personalized Guidance, P64).

Interestingly, when participants correctly identified AI responses, they acknowledged that AI could simulate sympathy or concern, albeit with limitations. Though lacking personal depth, one student remarked that an AI response “did express sympathy regardless of how lackluster it seemed” (Personal Connection, P105), and another noted that “it could have been more to act on, like meeting up, but they did provide other options for help” (Personal Connection, P97). In contrast, accurately identified human responses were seen as invested

in student success, but not necessarily emotionally expressive: “polite and invested in the student’s success but not super emotionally supportive” (Personal Connection, P4).

5 Discussion

Our study investigated whether students’ ratings of helpfulness and caring predicted their accuracy in identifying the source of mentorship responses (human vs. AI) across different domains. We found that higher ratings of both helpfulness and caring were associated with lower accuracy in source identification. This finding suggests that students equate warmth and supportiveness with human authorship, activating perceptual biases that limit their recognition of the potential support AI could provide. The more an AI-generated response resembles a human response, the less likely students are to recognize that it was written by AI.

These findings have critical implications for the integration of AI in mentorship contexts. Students may undervalue or distrust AI-generated guidance when it contradicts their assumptions about what AI can do. Even when AI performs well (e.g., offering emotional validation, supportive tone, or actionable advice), it is often misattributed or dismissed if its source is known. This bias may undermine trust in AI, particularly when relational authenticity is expected. While we previously attributed low source accuracy in the personal domain to students’ unfamiliarity with using AI in such contexts, our mixed-methods findings offer a more refined explanation. Students appear to use emotional tone as a heuristic for authorship, interpreting both overly emotional and insufficiently emotional responses as AI-generated. In contrast, they associate human-authored responses with balanced, relationally calibrated communication. Thus, when a response deviates from this midpoint, either too cold or too warm, it violates expectations and is more likely to be attributed to AI.

Furthermore, these insights raise several considerations for AI-supported mentorship in higher education. First, students often expect relational support from humans and assume that AI has its limitations in providing it. This calls for educational interventions to demystify what AI is capable of, especially in relational contexts, and promote informed AI literacy. Second, institutions and AI-designers may develop AI systems tailored to the institution’s specific policies, resources, and cul-

tural context, similar to how businesses invest in AI chatbots. This solution could help AI resources feel more as an ethical and trustworthy resource. Third, AI could handle surface-level information requests or initial support, freeing human mentors from trivial tasks or duties to provide deeper and emotionally nuanced engagement. Rather than replacing human mentorship, AI could enhance it when used as a complementary resource. Lastly, just as students misperceive AI as human based on warmth, they also misperceive humans as AI when their tone is rigid, detached, or overly formal. Institutions might consider providing training for their faculty and staff members on relational communication strategies, especially in email or digital interactions, to ensure that students are supported, even in brief exchanges.

Our study is not without limitations. First, the explanatory power of our models was weak (Model 1 T_{jur} ’s $R^2 = .015$, Model 2 $R^2 = .017$). These values suggest that, while the predictors were statistically significant, they account for only a small proportion of the variance in the accuracy of the source identification. Future research should replicate these findings using a larger sample size, a greater variety of stimuli, and more diverse educational contexts to improve generalizability. Second, the participant pool was limited to undergraduate students from a single liberal arts college in the northeastern United States. As such, the findings may reflect institution-specific dynamics and should be interpreted as exploratory or case-based. Expanding this research to include participants from multiple institutions and institutional types (e.g., community colleges, large public universities) would provide a more comprehensive understanding of students’ perceptions of AI and human mentorship. Third, although all faculty responses in this study were entirely human-authored without any AI assistance, it is possible that participants may have assumed that the faculty used AI tools to help craft their replies. Furthermore, the format in which responses were presented, modeled after email-based communication, may have influenced how participants perceived both the content and the source. AI responses framed as email replies may have appeared more human-like than if they were delivered through a chatbot or system-generated interface. This framing could have unintentionally blurred distinctions between human and AI authorship. Future research should investigate how different presentation formats (e.g., email, chatbot,

forum post) shape students' assumptions about authorship and credibility, and compare perceptions of human-authored, AI-assisted human-authored, and AI-authored mentorship responses across these contexts. Lastly, future studies would benefit from examining demographic variables such as race, ethnicity, and gender. Understanding how students from diverse backgrounds interpret and engage with AI-generated versus human mentorship may yield important insights, particularly as institutions strive to promote equity and culturally responsive mentorship practices.

Perceptual bias is not simply a barrier to AI adoption, but is a lens through which students interpret support and relational intent. Our results show that emotionally resonant, helpful responses are often mistaken for humans regardless of authorship, while detached or impersonal responses are perceived as AI. Emotional tone and personalization appeared to be more influential than the actual source in shaping students' evaluations. Yet, these biases are not fixed. As students gain more exposure to AI and as these tools become more embedded in academic settings, their ability to discern source and engage with AI more responsibly and meaningfully may improve. Our findings and recommendations provide a reflection of deeper sociocultural expectations about relational care, authenticity, and the boundaries between human and machine. The future of AI mentorship depends not just on technical capability, but on thoughtful, human-centered design that attends to the cognitive and relational dynamics in higher education settings.

Acknowledgments

We thank our participants for their time in taking part in this study. We also extend many thanks to faculty members who contributed their time to generating stimuli. Additionally, we acknowledge ChatGPT (version 3.5) for its role in stimulus generation.

References

- Shaki Asgari and Frederick Carter. 2016. [Peer mentors can improve academic performance: A quasi-experimental study of peer mentorship in introductory courses](#). *Teaching of Psychology*, 43(2):131–135.
- Dominic Baidoo-Anu and Linda Owusu Ansah. 2023. [Education in the era of generative artificial intelligence \(AI\): Understanding the potential benefits of ChatGPT in promoting teaching and learning](#). *SSRN Electronic Journal*.

- Nancy Budwig, Jimin Lee, and Raquel Jorge Fernandes. 2023. [A developmental and sociocultural approach to the transition from high school to college: The importance of understanding student meaning-making](#). *Human Arenas*.
- Joshua Dempere, Kiran Modugu, Ahmed Hesham, and Lakshmi K. Ramasamy. 2023. [The impact of ChatGPT on higher education](#). *Frontiers in Education*, 8.
- Gökhan Dost and Laura Mazzoli Smith. 2023. [Understanding higher education students' sense of belonging: A qualitative meta-ethnographic analysis](#). *Journal of Further and Higher Education*, 47(6):822–849.
- Satu Elo and Helvi Kyngäs. 2008. [The qualitative content analysis process](#). *Journal of Advanced Nursing*, 62(1):107–115.
- Tom Farrelly and Nigel Baker. 2023. [Generative artificial intelligence: Implications and considerations for higher education practice](#). *Education Sciences*, 13(11):1109.
- Jane Forman and Laura Damschroder. 2008. [Qualitative content analysis](#). In Laura Damschroder, editor, *Empirical methods for bioethics: A primer*, volume 11 of *Advances in Bioethics*, pages 39–62. Emerald/Elsevier.
- Mark Glickman and Tali Sharot. 2025. [How human–AI feedback loops alter human perceptual, emotional, and social judgements](#). *Nature Human Behaviour*, 9:345–359.
- Diala Hammoudi Halat, Amir Soltani, Rachid Dalli, Lujain Alsarraj, and Ahmed Malki. 2023. [Understanding and fostering mental health and well-being among university faculty: A narrative review](#). *Journal of Clinical Medicine*, 12(13):4425.
- Mary Jacobi. 1991. [Mentoring and undergraduate academic success: A literature review](#). *Review of Educational Research*, 61(4):505–532.
- Mohd Javaid, Abid Haleem, Rajiv P. Singh, Shujaat Khan, and I. H. Khan. 2023. [Unlocking the opportunities through ChatGPT tool towards ameliorating the education system](#). *BenchCouncil Transactions on Benchmarks, Standards and Evaluations*, 3(2):100115.
- Nils Köbis and Christian Mehner. 2021. [Ethical questions raised by AI-supported mentoring in higher education](#). *Frontiers in Artificial Intelligence*, 4:624050.
- Helvi Kyngäs. 2020. [Inductive content analysis](#). In *The application of content analysis in nursing science research*, pages 13–21. Springer.

- Hanh Le, Yulong Shen, Zhaoyuan Li, Minghui Xia, Linying Tang, Xiaoyu Li, Jia Jia, Qiyun Wang, Dragan Gašević, and Yubo Fan. 2025. [Breaking human dominance: Investigating learners' preferences for learning feedback from generative ai and human tutors](#). *British Journal of Educational Technology*, 56:1758–1783.
- Jimin Lee and Alena G. Esposito. 2025a. [ChatGPT or human mentors? student perceptions of technology acceptance and use and the future of mentorship in higher education](#). *Education Sciences*, 15(6):746.
- Jimin Lee and Alena G. Esposito. 2025b. A comparative study between college student perception of an AI learning tool and human mentoring responses. Manuscript submitted for publication.
- Jimin Lee, Si Wang, and Alena G. Esposito. 2025. Content analysis of college student perceptions: Mental health-related support from generative AI versus faculty mentors. *American Journal of Qualitative Research*. In press.
- Douglas Mayo and Benjamin Le. 2021. [Perceived discrimination and mental health in college students: A serial indirect effects model of mentoring support and academic self-concept](#). *Journal of American College Health*, 71(4):1184–1195.
- Marta Montenegro-Rueda, José Fernández-Cerero, José M. Fernández-Batanero, and Eloy López-Meneses. 2023. [Impact of the implementation of ChatGPT in education: A systematic review](#). *Computers*, 12(8):153.
- Wendy Nuis, Mien Segers, and Simon Beausaert. 2023. [Conceptualizing mentoring in higher education: A systematic literature review](#). *Educational Research Review*, 41:100565.
- OpenAI. 2024a. Chatgpt. <https://chat.openai.com/>. Conversational AI system; accessed March 2024.
- OpenAI. 2024b. Gpt-3.5 (gpt-3.5-turbo) model card and usage. <https://platform.openai.com/docs/models#gpt-3-5>. Model version used in this study; accessed March 2024.
- R Core Team. 2024. *R: A language and environment for statistical computing*. R Foundation for Statistical Computing, Vienna, Austria.
- Jean E. Rhodes, Jean Baldwin Grossman, and Nancy L. Resch. 2000. [Agents of change: Pathways through which mentoring relationships influence adolescents' academic adjustment](#). *Child Development*, 71(6):1662–1671.
- Michael Schönberger. 2023. [ChatGPT in higher education: The good, the bad, and the university](#). In *Proceedings of the 9th International Conference on Higher Education Advances (HEAd'23)*, pages 331–338.
- Danya F. Vears and Lynn Gillam. 2022. [Inductive content analysis: A guide for beginning qualitative researchers](#). *Focus on Health Professional Education: A Multi-Professional Journal*, 23(1):111–127.
- Viswanath Venkatesh, Michael G. Morris, Gordon B. Davis, and Fred D. Davis. 2003. [User acceptance of information technology: Toward a unified view](#). *MIS Quarterly*, 27(3):425–478.
- Xueming Wang, Shweta Hegde, Chiho Son, Bethany Keller, Alexander Smith, and Farzan Sasangohar. 2020. [Investigating mental health of US college students during the COVID-19 pandemic: Cross-sectional survey study](#). *Journal of Medical Internet Research*, 22(9):e22817.